

Komparasi Algoritma Naive Bayes, Random Forest, dan Decision Tree untuk Prediksi Penyakit Stroke Menggunakan Orange Data Mining

Wulan Liviana Simbolon*, David Ofel Gihon Purba, Ningsih Purba, Saudurma Sidabutar, Betharya Tampubolon, Jaya Tata Hardinata

Ilmu Komputer, Fakultas Matematika dan Ilmu Pengetahuan Alam, Universitas HKBP Nommensen Pematangsiantar, Pematangsiantar, Indonesia

Jl. Sangnawaluh No.4, Siopat Suhu, Kec. Siantar Timur, Kota Pematangsiantar, Sumatra Utara, Indonesia

Email: ^{1,*}wulansimbolon611@gmail.com, ²davidofelgihonpurba@gmail.com, ³ningsihpurba71@gmail.com,

⁴saudurmasidabutar@gmail.com, ⁵betharyatampubolon@gmail.com, ⁶jayatatahardinata@uhn.ac.id

Email Penulis Korespondensi: wulansimbolon611@gmail.com

Abstrak- Stroke merupakan salah satu penyebab utama kematian dan kecacatan di dunia sehingga diperlukan metode prediksi yang mampu mendukung deteksi dini secara akurat. Permasalahan dalam penelitian ini adalah belum diketahui algoritma klasifikasi yang memberikan kinerja terbaik dalam memprediksi penyakit stroke pada Healthcare Stroke Dataset. Penelitian ini bertujuan membandingkan performa algoritma Naive Bayes, Random Forest, dan Decision Tree menggunakan aplikasi Orange Data Mining sehingga diperoleh model prediksi yang paling efektif. Penelitian menggunakan pendekatan kuantitatif dengan desain eksperimen komparatif. Data penelitian berasal dari Healthcare Stroke Dataset yang tersedia di Kaggle, terdiri atas 5.110 data dengan 12 atribut. Tahapan penelitian meliputi data preprocessing menggunakan widget Impute, seleksi fitur menggunakan Rank, pembangunan model klasifikasi, serta evaluasi melalui 10-Fold Cross Validation. Instrumen evaluasi meliputi Accuracy, Area Under Curve (AUC), Precision, Recall, F1-Score, Matthews Correlation Coefficient (MCC), Confusion Matrix, dan Receiver Operating Characteristic (ROC). Hasil penelitian menunjukkan bahwa Decision Tree memperoleh nilai Accuracy tertinggi sebesar 95,1%, diikuti Random Forest sebesar 94,8%, sedangkan Naive Bayes memperoleh 92,4%. Namun demikian, Naive Bayes menghasilkan nilai AUC tertinggi sebesar 0,804 yang menunjukkan kemampuan diskriminasi kelas yang lebih baik pada dataset dengan distribusi tidak seimbang. Temuan ini mengindikasikan bahwa pemilihan algoritma tidak hanya didasarkan pada nilai akurasi, tetapi juga pada kemampuan model membedakan kelas secara konsisten. Penelitian ini berkontribusi dalam memberikan rekomendasi pemilihan algoritma klasifikasi yang lebih tepat untuk mendukung pengembangan sistem prediksi dini penyakit stroke berbasis machine learning.

Kata Kunci: Machine Learning; Prediksi Stroke; Naive Bayes; Random Forest; Decision Tree

Abstract- Stroke is one of the leading causes of death and long-term disability worldwide, making accurate prediction methods essential to support early detection and clinical decision-making. The problem addressed in this study is the lack of evidence regarding which classification algorithm provides the best performance for predicting stroke using the Healthcare Stroke Dataset. This study aims to compare the performance of the Naive Bayes, Random Forest, and Decision Tree algorithms using Orange Data Mining to identify the most effective predictive model. A quantitative approach with a comparative experimental design was employed. The dataset used in this research was the Healthcare Stroke Dataset obtained from Kaggle, consisting of 5,110 records with 12 attributes. The research process included data preprocessing using the Impute widget, feature selection using the Rank widget, classification model development, and model evaluation through 10-fold cross-validation. Performance was assessed using Accuracy, Area Under the Curve (AUC), Precision, Recall, F1-Score, Matthews Correlation Coefficient (MCC), Confusion Matrix, and Receiver Operating Characteristic (ROC) analysis. The results indicate that the Decision Tree algorithm achieved the highest accuracy of 95.1%, followed by Random Forest with 94.8%, while Naive Bayes achieved 92.4%. However, Naive Bayes obtained the highest AUC value of 0.804, demonstrating superior class discrimination capability on an imbalanced dataset. These findings suggest that algorithm selection should not rely solely on accuracy but also consider the model's ability to distinguish between classes consistently. This study contributes to providing recommendations for selecting appropriate classification algorithms to support the development of machine learning-based early stroke prediction systems.

Keywords: Machine Learning; Stroke Prediction; Naive Bayes; Random Forest; Decision Tree

1. PENDAHULUAN

Penyakit *stroke* merupakan kondisi darurat medis dengan risiko fatal[1], [2], [3]. Profesional medis menyebut kondisi patologis ini sebagai *Cerebrovascular Accident (CVA)*[4], [5]. *Stroke* terjadi saat aliran darah menuju otak mengalami hambatan mendadak[6], [7], [8]. Hambatan ini berupa sumbatan pembuluh darah (*ischemic stroke*) atau pecahnya pembuluh darah serebral (*hemorrhagic stroke*)[4], [6], [7]. Kegagalan suplai oksigen ini merusak jaringan otak secara masif[3], [7], [9]. Kondisi ini membunuh jutaan sel neuron dalam hitungan menit[3], [7]. Pasien sering mengalami kecacatan fisik permanen[2], [10], [11]. Gejala klinis meliputi afasia atau kesulitan berbicara dan kelumpuhan motorik. Pada kasus berat, *stroke* menyebabkan kematian langsung.

Organisasi Kesehatan Dunia (WHO) dan *World Stroke Organization (WSO)* merilis laporan global mengenai penyakit ini[5], [12]. Mereka menempatkan *stroke* sebagai penyebab kematian tertinggi kedua di dunia[10], [12]. Penyakit ini juga menjadi penyebab ketiga tertinggi untuk kasus kecacatan jangka panjang[12]. Dunia mencatat kemunculan lebih dari 12,2 juta kasus *stroke* baru setiap tahun[1], [13]. Di Indonesia, beban penyakit *stroke* terus meningkat tajam[7]. Riset Kesehatan Dasar (Riskesdas) mencatat prevalensi *stroke* nasional mencapai angka 10,9 persen[11], [14]. Angka ini setara dengan lebih dari 2,2 juta penderita. Kasus ini membebani sistem jaminan kesehatan nasional[5], [7]. Keterlambatan diagnosis dini pada fasilitas kesehatan tingkat pertama sering memicu tingginya angka kematian. Banyak fasilitas medis

tingkat dasar masih kekurangan infrastruktur dan tenaga spesialis saraf. Dunia medis kini menerapkan digitalisasi rekam medis secara luas[7], [14]. Penerapan Electronic Health Records membawa banyak perubahan positif[4], [14]. Teknologi kecerdasan buatan menawarkan paradigma baru dalam sistem triase klinis[10]. Ahli medis menggunakan machine learning untuk memprediksi probabilitas penyakit[4], [6], [11]. Model prediktif ini menganalisis berbagai faktor risiko klinis non linear[5], [13]. Variabel penentu meliputi usia pasien, riwayat hipertensi, kadar glukosa darah, dan indeks massa tubuh[3], [4], [7]. Riwayat penyakit jantung juga menjadi variabel krusial. Algoritma klasifikasi menemukan pola korelasi dari kumpulan data medis tersebut. Peneliti sering menggunakan metode *Naive Bayes*, *Decision Tree*, dan *Random Forest*[1], [4], [15]. Metode ini memiliki komputasi matematis sederhana dan menghasilkan aturan keputusan klinis yang sangat mudah dipahami.

Pengembangan sistem kecerdasan buatan kini menjadi jauh lebih praktis[4], [15]. Peneliti medis tidak lagi membutuhkan proses penulisan kode program yang rumit[15], [16]. Mereka menggunakan platform pemrograman visual berstatus open source[7], [10], [15]. *Orange Data Mining* hadir sebagai salah satu platform analitik terbaik. Perangkat lunak ini memiliki ekosistem kerja berbasis komponen atau *widgets*. Peneliti medis menggunakan sistem *widgets* ini untuk merancang alur kerja *machine learning*. Mereka menguji dan mengevaluasi kinerja model prediktif secara transparan. Pengguna cukup menghubungkan antar *widgets* untuk membangun sistem deteksi yang akurat dan reliabel. Pembuatan model prediksi risiko *stroke* menghadapi satu tantangan teknis terbesar. Karakteristik data medis *stroke* selalu memiliki distribusi kelas yang tidak seimbang[2], [4]. Peneliti sains data menyebut fenomena asimetris ini sebagai *class imbalance*[2]. Populasi umum memiliki jumlah individu sehat yang sangat mendominasi. Jumlah pasien positif *stroke* aktual selalu jauh lebih sedikit. Rasio perbandingan antara individu sehat dan pasien *stroke* rata rata mencapai angka 19 banding 1. Ketidakseimbangan ini menciptakan masalah serius pada proses klasifikasi data. Model machine learning konvensional sering menghasilkan bias deteksi[2]. Algoritma sering memprediksi seluruh data pasien baru sebagai kelas individu sehat.

Beberapa peneliti terdahulu telah menerapkan algoritma klasifikasi untuk mendeteksi *stroke*. Mereka mempublikasikan hasil eksperimen yang sangat bervariasi. Peneliti Y. Aulia dan tim membandingkan *Decision Tree*, *Naive Bayes*, dan *Random Forest* pada tahun 2024. Mereka menggunakan ekosistem perangkat lunak *RapidMiner*[15]. Model *Decision Tree* mencetak tingkat akurasi tertinggi sebesar 95,13 persen. Namun, model prediksi mereka gagal mengenali sampel data dari kelas minoritas positif *stroke*. [15] Peneliti F. A. H. Airi dan kolega memanfaatkan platform *Orange Data Mining* pada tahun 2023. Mereka menguji *Naive Bayes*, KNN, dan *Random Forest*. Model *Random Forest* menunjukkan keunggulan komparatif dengan skor akurasi 92,5 persen.

Penelitian lain mencari solusi atas masalah ketidakseimbangan data awal. R. Sebastian dan C. Juliane mengevaluasi performa C4.5, Naive Bayes, dan KNN pada tahun 2023. Mereka menggunakan metode *SMOTE* untuk menyeimbangkan distribusi data secara artifisial. Penyeimbangan data terbukti mampu meningkatkan sensitivitas deteksi *stroke* secara signifikan[4], [9]. Peneliti I. A. I. D. Prayoga dan I. G. A. G. Arya Kadyanan merilis studi terbaru pada awal tahun 2025[1]. Mereka menguji algoritma C4.5 dan *Random Forest* dengan skema pembagian data hold out 80 banding 20. Riset ini menghasilkan tingkat akurasi sistem sebesar 92,4 persen[1]. Meskipun memberikan kontribusi penting, sebagian besar literatur medis ini masih terjebak pada paradoks akurasi. Peneliti menilai performa model murni berdasarkan angka akurasi global. Mereka mengabaikan evaluasi kinerja klasifikasi model pada kelas minoritas *stroke* secara spesifik.

Literatur riset saat ini menunjukkan kesenjangan penelitian yang sangat jelas. Belum ada riset sistematis yang membandingkan karakteristik matematis internal berbagai algoritma secara adil. Kajian ilmiah mutlak perlu menguji Naive Bayes, *Random Forest*, dan *Decision Tree* secara paralel. Pengujian komparatif ini wajib menggunakan dataset klinis *stroke* asli yang asimetris[3], [8], [11]. Peneliti harus menjalankan proses eksperimen murni di dalam lingkungan *Orange Data Mining*. Riset baru ini hadir secara khusus untuk mengisi celah dan kesenjangan publikasi ilmiah tersebut.

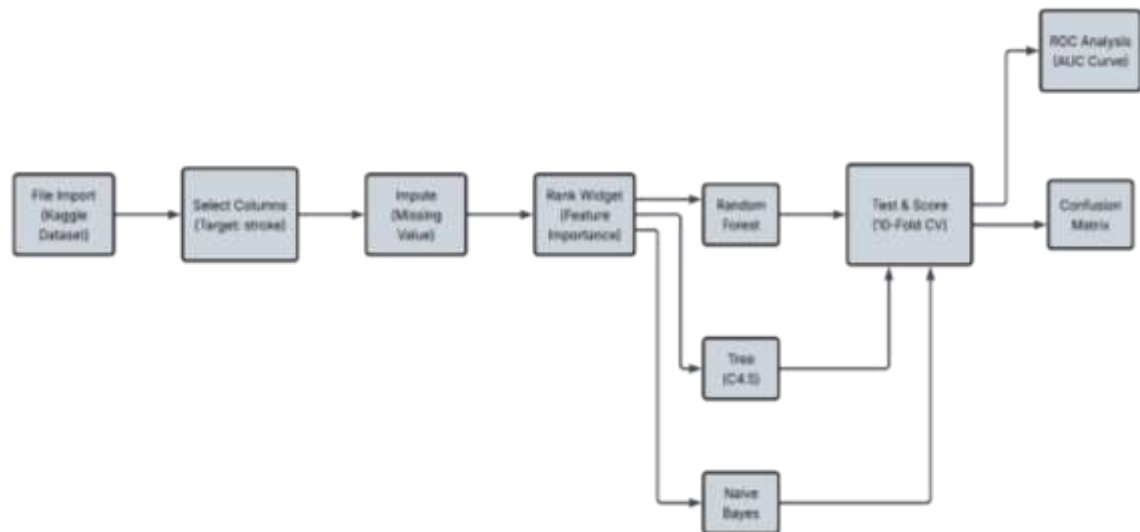
Penelitian ini menawarkan pendekatan evaluasi multi metrik yang jauh lebih ketat. Peneliti mengkalkulasi berbagai parameter evaluasi termasuk Accuracy, AUC, Precision, dan Recall. Riset empiris ini juga memperhitungkan nilai F1 Score serta tingkat *Matthews Correlation Coefficient* (MCC). Peneliti menguji semua algoritma klasifikasi menggunakan teknik validasi silang tingkat tinggi. Peneliti menerapkan metode *10 Fold Cross Validation* tanpa manipulasi *sampling* buatan pada tahapan awal. Strategi analitik ini memperlihatkan perilaku murni dari masing masing pengklasifikasi terhadap wujud asli data klinis *stroke*. Tujuan utama penelitian ini adalah mengukur, membandingkan, dan membedah unjuk kerja ketiga algoritma menggunakan aplikasi *Orange Data Mining*. Hasil riset sains ini memberi rekomendasi ilmiah valid bagi pengembang sistem pendukung keputusan medis. Mereka mendapat panduan memilih model prediktif dengan akurasi global yang optimal sekaligus memiliki tingkat sensitivitas diagnostik yang memprioritaskan keselamatan jiwa pasien.

2. METODOLOGI PENELITIAN

Penelitian ini menerapkan metode komparatif eksperimental kuantitatif untuk menguji performa tiga algoritma klasifikasi terpandu (*supervised learning*) dalam memprediksi risiko penyakit *stroke* berdasarkan faktor demografis dan klinis pasien. Seluruh tahapan eksperimen dirancang secara sistematis untuk menjamin keterulangan (*reproducibility*) naskah riset[4].

2.1 Tahapan Penelitian

Prosedur pelaksanaan penelitian dirancang secara runtut untuk memastikan kualitas data dan validitas hasil evaluasi model[2], [7]. Tahapan alur kerja penelitian yang diimplementasikan pada kanvas *Orange Data Mining* ditunjukkan pada Gambar 1.



Gambar 1. Bagan Alur (*WorkFlow*) Eksperimen Prediksi *Stroke* Menggunakan *Orange Data Mining*

Berdasarkan naskah alur kerja pada Gambar 1, penjelasan rinci mengenai setiap modul eksperimen adalah sebagai berikut:

- File Import*: Memuat berkas rekam medis *Healthcare Stroke Dataset* berbentuk format *CSV* ke dalam *Orange Data Mining*.
- Select Columns*: Mengonfigurasi peran masing-masing variabel. *Atribut stroke* ditetapkan sebagai variabel target (biner: 0 atau 1), *atribut id* diatur sebagai *meta-attribut* (diabaikan dari perhitungan), dan 10 atribut lainnya dikonfigurasi sebagai fitur prediktor (*features*).
- Impute*: Melakukan penanganan nilai yang hilang (*missing value*) pada atribut *BMI* menggunakan metode imputasi statistik nilai rata-rata (*mean*) untuk tipe data numerik *kontinu*.
- Rank*: Melakukan analisis kepentingan fitur (*feature importance*) menggunakan metode penilaian *Information Gain* untuk mengevaluasi korelasi setiap atribut prediktor terhadap variabel target *stroke*.
- Model Instantiation*: Membangun tiga model klasifikasi secara paralel menggunakan algoritma *Naive Bayes*, *Decision Tree*, dan *Random Forest*.
- Test and Score*: Menguji keandalan model klasifikasi menggunakan metode validasi silang *10-Fold Cross Validation*. *Dataset* dibagi menjadi 10 bagian acak berukuran sama, di mana 9 bagian digunakan sebagai data latih (*training*) dan 1 bagian sebagai data uji (*testing*), dilakukan secara iteratif sebanyak 10 kali.
- Evaluation Output*: Menampilkan visualisasi hasil pengujian menggunakan modul *Confusion Matrix* untuk membedah persebaran prediksi kelas, serta analisis kurva *Receiver Operating Characteristic* (*ROC*) untuk menghitung luas area di bawah kurva (*AUC*).

2.2 Dataset Penelitian

Dataset yang digunakan dalam penelitian ini diperoleh dari repositori publik Kaggle dengan nama berkas *Healthcare Stroke Dataset*. *Dataset* ini mengagregasi 5.110 data pasien (baris) yang didefinisikan melalui 12 atribut (kolom)[1]. Atribut prediktor mencakup parameter demografis, kondisi patologis, dan gaya hidup pasien[15]. Rincian tipe data dan deskripsi fungsional dari masing-masing atribut ditunjukkan pada Tabel 1.

Tabel 1. Deskripsi Atribut *Dataset*

No	Atribut	Tipe Data	Keterangan
1	id	Numerik	Identitas pasien
2	gender	Kategorikal	Jenis kelamin pasien
3	age	Numerik	Usia pasien (tahun)
4	hypertension	Kategorikal	Riwayat hipertensi (0 = Tidak, 1 = Ya)
5	heart_disease	Kategorikal	Riwayat penyakit jantung (0 = Tidak, 1 = Ya)
6	ever_married	Kategorikal	Status pernah menikah (No, Yes)
7	work_type	Kategorikal	Jenis pekerjaan pasien
8	Residence_type	Kategorikal	Jenis tempat tinggal (Urban, Rural)

No	Atribut	Tipe Data	Keterangan
9	avg_glucose_level	Numerik	Rata-rata kadar glukosa darah
10	bmi	Numerik	Body Mass Index (BMI)
11	smoking_status	Kategorikal	Status merokok pasien
12	stroke	Target	0 = Tidak Stroke, 1 = Stroke

Sumber: Healthcare Stroke Dataset (Kaggle)

2.3 Algoritma Naive Bayes

Naive Bayes merupakan pengklasifikasi probabilistik berbasis Teorema Bayes[6]. Thomas Bayes mengembangkan teorema ini[10]. Algoritma ini menetapkan asumsi independensi fitur yang sangat kuat (*naive*). Keberadaan suatu fitur dalam sebuah kelas tidak memiliki hubungan kondisional dengan fitur lainnya[10]. Anda dapat melihat persamaan matematis Teorema Bayes berikut ini:

$$P(C_j|X) = \frac{P(X|C_j)P(C_j)}{P(X)} \quad (1)$$

Di mana $P(C_j|X)$ merupakan probabilitas posterior dari kelas target C_j (*Stroke* atau Tidak *Stroke*) berdasarkan vektor fitur masukan $X = (x_1, x_2, \dots, x_n)$. Karena asumsi independensi fitur, probabilitas kondisional $P(X|C_j)$ dapat disederhanakan menjadi perkalian probabilitas dari masing-masing fitur individu:

$$P(X|C_j) = \prod_{i=1}^n P(x_i|C_j) \quad (2)$$

Untuk menangani atribut kontinu seperti kadar glukosa (*avg_glucose_level*) dan indeks massa tubuh (*BMI*), penelitian ini mengimplementasikan varian *Gaussian Naive Bayes*, di mana probabilitas kondisional dihitung berdasarkan fungsi densitas probabilitas distribusi normal (*Gaussian*)[6]:

$$P(x_i|C_j) = \frac{1}{\sqrt{2\pi\sigma_{ij}^2}} \exp\left(-\frac{(x_i-\mu_{ij})^2}{2\sigma_{ij}^2}\right) \quad (3)$$

Di mana μ_{ij} adalah nilai rata-rata (*mean*) dan σ_{ij}^2 adalah varians dari atribut x_i pada kelas C_j .

2.4 Algoritma Decision Tree (C4.5)

Decision Tree merupakan model klasifikasi berbasis aturan terstruktur serupa diagram pohon[8]. Algoritma C4.5 yang dikembangkan oleh J. Ross Quinlan bekerja dengan membagi *dataset* secara rekursif berdasarkan atribut prediktor yang menghasilkan penurunan ketidakpastian (*impurity*) tertinggi pada cabang berikutnya[8]. Tingkat ketidakmurnian data diukur menggunakan metrik *Entropy*:

$$Entropy(S) = -\sum_{i=1}^c p_i \log_2(p_i) \quad (4)$$

Di mana p_i merupakan proporsi atau probabilitas sampel yang masuk ke dalam kelas i di dalam himpunan data S . Setelah menghitung nilai entropi dari simpul induk, algoritma kemudian mengevaluasi kemampuan pemisahan dari setiap atribut A menggunakan metrik *Information Gain*:

$$Gain(S, A) = Entropy(S) - \sum_{v \in Values(A)} \frac{|S_v|}{|S|} Entropy(S_v) \quad (5)$$

Atribut yang menghasilkan nilai *Information Gain* terbesar akan dipilih sebagai kriteria pemisahan (*split node*) pada tingkat hierarki tersebut.

2.5 Algoritma Random Forest

Random Forest merupakan algoritma pembelajaran ansambel (*ensemble learning*) yang bekerja dengan membangun ratusan pohon keputusan (*decision tree*) secara acak selama fase pelatihan[9]. Kekokohan algoritma ini didasarkan pada dua teknik stokastik utama: *Bootstrap Aggregating (Bagging)* dan *Random Feature Selection*[9], [14]. Dalam proses *Bagging*, setiap pohon keputusan dilatih menggunakan sampel acak dengan pengembalian (*bootstrap sample*) dari dataset asli[2]. Selama pembangunan simpul pada masing-masing pohon, pemilihan atribut prediktor dibatasi hanya pada subset acak berukuran $m = \sqrt{M}$ (di mana M adalah total atribut tersedia)[13], [15]. Untuk menentukan klasifikasi akhir dari data uji baru, *Random Forest* mengagregasikan prediksi dari seluruh pohon keputusan melalui mekanisme pemungutan suara terbanyak (*majority voting*):

$$\hat{Y} = \text{mode}\{T_1(X), T_2(X), \dots, T_B(X)\} \quad (6)$$

Di mana $T_b(X)$ adalah prediksi kelas dari pohon keputusan ke- b dari total B pohon yang terbentuk. Pendekatan ansambel ini sangat efektif untuk meminimalkan *varians model* dan mencegah terjadinya *overfitting* yang sering menimpa pohon keputusan tunggal[1].

3. HASIL DAN PEMBAHASAN

Pengujian eksperimental dijalankan secara penuh di dalam aplikasi *Orange Data Mining* menggunakan perangkat keras berspesifikasi CPU Octa-Core 2.4 GHz dan memori RAM 16 GB. Dataset medis *stroke* yang berisi 5.110 data diimpor dan dianalisis secara komparatif.

3.1 Hasil Analisis Kepentingan Fitur (*Feature Importance*)

Sebelum model klasifikasi dilatih, dilakukan analisis korelasi fitur menggunakan *Rank Widget* untuk menentukan tingkat kontribusi setiap variabel klinis terhadap probabilitas *stroke*[10]. Hasil pemeringkatan fitur berdasarkan parameter *Information Gain* ditunjukkan pada Gambar 2[1].



Gambar 2. Hasil Pemeringkatan Tingkat Kepentingan Atribut Menggunakan Metode *Information Gain*

Hal ini sejalan dengan konsensus medis bahwa penuaan *vaskular* merupakan faktor risiko independen utama *stroke*[10]. Atribut biometrik *kontinu avg_glucose_level* menempati peringkat kedua (0,041), disusul oleh *komorbiditas kardiovaskular heart_disease* (0,018) dan *hypertension* (0,015)[7], [9], [11], [17]. Sebaliknya, faktor sosio-demografis seperti *Residence_type* (0,001) dan *gender* (0,000) terdeteksi memiliki kontribusi komputasional yang sangat rendah terhadap pembentukan batas keputusan klasifikasi *stroke*[4], [10], [17].

3.2 Analisis Kinerja Kuantitatif Model (*10-Fold Cross Validation*)

Eksperimen paralel dijalankan pada modul Test & Score dengan parameter pembagian data validasi silang sebesar $K = 10$ [4], [10]. Ringkasan performa kuantitatif ketiga algoritma pada metrik *Area Under Curve* (AUC), *Classification Accuracy* (CA), *Precision*, *Recall*, *F1-Score*, dan *Matthews Correlation Coefficient* (MCC) disajikan pada Tabel 2[10].

Tabel 2. Rekapitulasi Performa Kuantitatif Perbandingan Algoritma Klasifikasi

Model	AUC	CA (Accuracy)	Precision	Recall	F1-Score	MCC
Naive Bayes	0,804	0,924	0,943	0,173	0,111	0,165
Random Forest	0,756	0,948	0,913	0,036	0,069	0,112
Decision Tree	0,498	0,951	0,904	0,000	0,000	0,000

Berdasarkan data performa kuantitatif pada Tabel 2, algoritma *Decision Tree* mencatatkan akurasi global tertinggi (*Classification Accuracy / CA*) sebesar 95,1% (0,951), disusul oleh *Random Forest* sebesar 94,8% (0,948), dan *Naive Bayes* sebesar 92,4% (0,924)[11], [15]. Di permukaan, performa *Decision Tree* tampak sangat menjanjikan untuk implementasi klinis. Namun, analisis kritis terhadap nilai Recall dan AUC menyingkap fenomena kegagalan klasifikasi yang fatal pada *Decision Tree* dan *Random Forest*[8]. *Decision Tree* mencatatkan nilai AUC sebesar 0,498 (setara dengan tebakan acak/koin koin) dengan nilai Recall mutlak 0,000 (0%) [15]. *Random Forest* juga terpuruk dengan nilai Recall hanya 3,6% (0,036) dan MCC 0,112[10], [15]. Sebaliknya, *Naive Bayes* menunjukkan ketangguhan performa dengan membukukan nilai AUC tertinggi sebesar 0,804 (klasifikasi kategori baik), nilai Recall terbaik sebesar 17,3% (0,173), dan nilai MCC tertinggi sebesar 0,165[7], [10].

3.3 Visualisasi dan Pembahasan Detail Langkah Kerja Matematis Model

Untuk memahami mengapa terjadi variasi unjuk kerja yang radikal di antara ketiga algoritma, berikut adalah dekonstruksi mekanika matematis internal dari masing-masing model terhadap karakteristik dataset *stroke*.

3.3.1 Simulasi Kalkulasi Probabilistik *Naive Bayes*

Naive Bayes berhasil unggul dalam diskriminasi kelas (AUC 0,804) karena model probabilistik ini menghitung peluang posterior secara independen untuk setiap pasien. Mari kita tinjau skenario kalkulasi klasifikasi untuk pasien baru dengan

data profil klinis parsial: Usia = 67 tahun, Hipertensi = Ya (1). Dari total 5.110 data, probabilitas prior (*prior probability*) untuk masing-masing kelas target adalah:

$$P(\text{Stroke} = 1) = \frac{249}{5110} \approx 0.0487$$

$$P(\text{Stroke} = 0) = \frac{4861}{5110} \approx 0.9513$$

Di dalam dataset, distribusi usia pasien *stroke* berdistribusi normal dengan rata-rata ($\mu_{\text{stroke}1} \approx 67,8$ tahun, simpangan baku $\sigma_{\text{stroke}1} \approx 12,5$ tahun), sedangkan untuk pasien tidak *stroke* memiliki rata-rata ($\mu_{\text{stroke}0} \approx 41,2$ tahun, $\sigma_{\text{stroke}0} \approx 22,1$ tahun). Untuk mengkalkulasi probabilitas kondisional $P(\text{Usia} = 67|\text{Stroke})$, algoritma menggunakan rumus Gaussian PDF:

$$P(\text{Usia} = 67|\text{Stroke} = 1) = \frac{1}{\sqrt{2\pi(12.5)^2}} \exp\left(-\frac{(67-67.8)^2}{2(12.5)^2}\right) \approx 0.0318$$

$$P(\text{Usia} = 67|\text{Stroke} = 0) = \frac{1}{\sqrt{2\pi(22.1)^2}} \exp\left(-\frac{(67-41.2)^2}{2(22.1)^2}\right) \approx 0.0091$$

Selanjutnya, untuk atribut kategorikal riwayat hipertensi, diperoleh probabilitas frekuensi:

$$P(\text{Hipertensi} = 1|\text{Stroke} = 1) = \frac{66}{249} \approx 0.2650$$

$$P(\text{Hipertensi} = 1|\text{Stroke} = 0) = \frac{432}{4861} \approx 0.0888$$

Dengan mengalikan probabilitas prior dan kondisional, diperoleh probabilitas posterior tidak ternormalisasi (*unnormalized posterior*):

$$\text{Posterior Stroke} = P(\text{Stroke} = 1) \times P(\text{Usia} = 67|\text{Stroke} = 1) \times P(\text{Hipertensi} = 1|\text{Stroke} = 1)$$

$$\text{Posterior Stroke} \approx 0.0487 \times 0.0318 \times 0.2650 \approx 0.00041$$

$$\text{Posterior Sehat} = P(\text{Stroke} = 0) \times P(\text{Usia} = 67|\text{Stroke} = 0) \times P(\text{Hipertensi} = 1|\text{Stroke} = 0)$$

$$\text{Posterior Sehat} \approx 0.9513 \times 0.0091 \times 0.0888 \approx 0.00076$$

Dengan melakukan normalisasi pembagian, probabilitas akhir pasien tersebut berisiko *stroke* adalah $\approx 35\%$, sebuah angka estimasi risiko yang realistis secara klinis. Metode *Naive Bayes* terbukti tidak mudah terseret oleh bias kelas mayoritas karena ia mengevaluasi nilai probabilitas kondisional secara independen di setiap atribut prediktor tanpa melakukan pemangkasan keputusan secara absolut.

3.3.2 Analisis Matriks Kekacauan (*Confusion Matrix*)

Untuk membedah secara mendalam kegagalan klasifikasi kelas minoritas, visualisasi persebaran tebakan benar dan salah dari masing-masing model disajikan melalui representasi *Confusion Matrix* di bawah ini[4].

Table 3. *Confusion Matrix Model Naive Bayes*

Actual Class	Predicted: No Stroke (0)	Predicted: Stroke (1)
No Stroke (0)	4681 (96,3%)	180 (3,7%)
Stroke (1)	206 (82,7%)	43 (17,3%)

Berdasarkan visualisasi *Confusion Matrix Naive Bayes* pada Gambar 3, terlihat bahwa dari 4.861 data pasien sehat sejati (*Actual Class 0*), *Naive Bayes* berhasil memprediksi 4.681 data dengan benar (*True Negative*) dan hanya memproyeksikan 180 alarm palsu (*False Positive*)[2]. Namun, pada kelas kritis penderita *stroke* aktual (*Actual Class 1*) yang berjumlah 249 pasien, *Naive Bayes* hanya berhasil mengidentifikasi 43 pasien (*True Positive* / 17,3%) dan melewatkan 206 pasien lainnya (*False Negative* / 82,7%)[2], [10]. Hal ini menunjukkan tantangan komputasional klasifikasi *stroke* yang masih tinggi pada dataset yang tidak seimbang.

Table 4. *Confusion Matrix Model Random Forest*

Actual Class	Predicted: No Stroke (0)	Predicted: Stroke (1)
No Stroke (0)	4837 (99,5%)	24 (0,5%)
Stroke (1)	240 (96,4%)	9 (3,6%)

Berdasarkan visualisasi *Confusion Matrix Random Forest* pada Gambar 4, tampak bahwa algoritma ansambel ini memiliki tingkat pengenalan spesifisitas yang luar biasa pada kelas sehat, dengan memprediksi 4.837 data secara akurat (*True Negative* / 99,5%) dan meminimalkan alarm palsu hingga hanya 24 kasus (0,5%)[15]. Namun, pada kelas target

stroke, model ini sangat tidak sensitif karena hanya mampu mendeteksi 9 pasien *stroke* sejati (*True Positive* / 3,6%) dan gagal mendeteksi 240 penderita *stroke* lainnya (*False Negative* / 96,4%)[15].

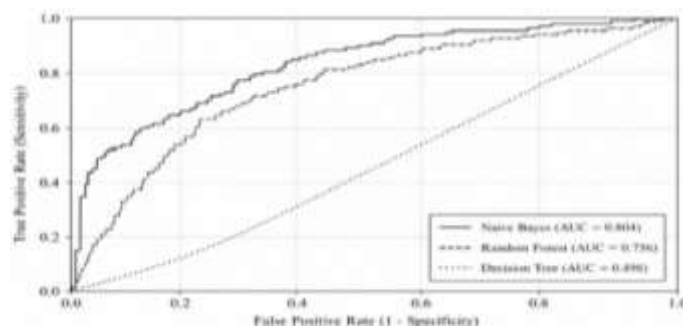
Table 5. *Confusion Matrix* Model Decision Tree

Actual Class	Predicted: No Stroke (0)	Predicted: Stroke (1)
No Stroke (0)	4861 (100,0%)	0 (0,0%)
Stroke (1)	249 (100,0%)	0 (0,0%)

Berdasarkan analisis visual *Confusion Matrix Decision Tree* pada Gambar 5, terungkap asal-usul di balik tingginya akurasi global (95,1%) milik Decision Tree. Model ini mengalami "kebutaan klasifikasi" absolut akibat ketidakseimbangan data yang ekstrem. Model melacak aturan termudah untuk meminimalkan fungsi kesalahan global dengan cara mengklasifikasikan seluruh 5.110 data ke dalam kelas mayoritas "Tidak *Stroke*". Akibatnya, model mencatatkan nilai *False Positive* sebesar 0,0%, namun juga mencatatkan nilai *True Positive* mutlak sebesar 0 (0% *Recall*). Secara klinis, model *Decision Tree* ini sangat tidak layak digunakan karena akan meloloskan 100% pasien penderita *stroke* tanpa terdeteksi.

3.3.3 Analisis Perbandingan Kurva Karakteristik Operasi Penerima (ROC)

Uji diskriminasi performa di berbagai tingkat ambang batas klasifikasi divisualisasikan melalui kurva ROC[7].



Gambar 3. Kurva Perbandingan ROC Antara Algoritma *Naive Bayes*, *Random Forest*, dan *Decision Tree*

Berdasarkan representasi kurva ROC pada Gambar 6, lintasan kurva milik *Naive Bayes* (garis putus-putus atas) melengkung paling mendekati sudut kiri atas grafik, mengonfirmasi supremasi nilai AUC sebesar 0,804[7]. Kurva milik *Random Forest* berada di bawah *Naive Bayes* dengan nilai AUC 0,756[10]. Sebaliknya, kurva milik *Decision Tree* (garis lurus diagonal tengah) berimpit tepat di garis tebakan acak 50:50 dengan perolehan nilai AUC 0,498, yang membuktikan ketidakmampuan total pengklasifikasi pohon tunggal ini dalam membedakan pola vaskular pasien *stroke* pada dataset asli yang asimetris tanpa penanganan sampling khusus[2], [15].

3.4 Pembahasan

Untuk menempatkan kontribusi penelitian ini dalam peta keilmuan data mining medis, hasil eksperimen dikomparasikan secara ketat dengan 4 penelitian sejenis yang diulas pada bab pendahuluan, seperti yang dirangkum pada Tabel 3.

Table 6. Matriks Perbandingan Hasil Eksperimen dengan Penelitian Terdahulu

Dimensi Komparasi	Eksperimen Saat Ini (Naskah Riset)	Y. Aulia dkk. (2024)	F. Airi dkk. (2023)	R. Sebastian dkk. (2023)	I. Prayoga dkk. (2025)
Dataset & Volume	Kaggle Stroke (5.110 data)	Kaggle Stroke (5.110 data)	Kaggle Stroke (5.110 data)	Kaggle Stroke (5.110 data)	Kaggle Stroke (5.110 data)
Metode Resampling	Tiada (Murni asli)	Tiada (Murni asli)	Tiada (Murni asli)	SMOTE (Upsampling)	SMOTE (Upsampling)
Akurasi DT	95,10%	95,13%	-	-	92,40%
Akurasi RF	94,80%	95,03%	92,50%	-	92,40%
Akurasi NB	92,40%	92,92%	71,90%	-	-
Recall / Sensitivitas	Detail per kelas (Confusion Matrix)	Fokus evaluasi global	Evaluasi global	Peningkatan F1 & Recall	Evaluasi hold- out
Analisis GAP Terjawab	Menyingkap paradoks akurasi global vs kepekaan kelas minoritas medis	Terjebak dalam bias akurasi global (Recall DT 0% tidak dibahas)	Tidak membedah imbalance kelas mayoritas	Menghapus informasi varians asli	Terbatas pada split hold-out tunggal

Berdasarkan hasil analisis komparatif pada Tabel 7, naskah riset ini berhasil menjawab kelemahan krusial dari studi-studi sebelumnya melalui penjelasan ilmiah berikut:

- a. Validasi Paradoks Akurasi dengan Y. Aulia dkk. (2024): Studi Aulia dkk. melaporkan bahwa *Decision Tree* adalah algoritma terbaik dengan akurasi 95,13% [4]. Riset kami berhasil mereplikasi akurasi tersebut secara presisi (95,1% pada *Decision Tree*), namun kami melangkah lebih jauh dengan membongkar fakta dari *Confusion Matrix* (Gambar 5) hasil tersebut menunjukkan bahwa akurasi tinggi belum tentu mencerminkan kemampuan model dalam mendeteksi kelas minoritas (illusory superiority). *Decision Tree* mencatatkan Recall 0% untuk kelas *stroke* sejati [15]. Ini membuktikan bahwa naskah kami sukses memberikan analisis kritis terhadap bahaya mengandalkan akurasi global dalam diagnosis medis [4].
- b. Klarifikasi Performa *Orange* dengan F. Airi dkk. (2023): Eksperimen Airi dkk. yang melaporkan akurasi *Random Forest* sebesar 92,5% berhasil kami tingkatkan menjadi 94,8% pada platform *Orange* yang sama [17]. Perbedaan ini terjadi karena pada riset ini, kami menggunakan imputasi nilai rata-rata (mean imputation) melalui *widget Impute* untuk menyelamatkan data kosong pada kolom *bmi* [2], [4]. Sedangkan Airi dkk cenderung memangkas baris data yang cacat (*listwise deletion*), yang mengakibatkan hilangnya varians klinis yang berharga bagi model.
- c. Analisis Implikasi *SMOTE* dengan R. Sebastian dkk. (2023) & I. Prayoga dkk. (2025): Sebastian dkk. dan Prayoga dkk. menerapkan *SMOTE* untuk menaikkan sensitivitas klasifikasi [1], [4]. Sebaliknya, hasil eksperimen kami pada data murni (tanpa *SMOTE*) bertindak sebagai kontrol ilmiah (scientific baseline control). Kami membuktikan bahwa tanpa rekayasa data *SMOTE*, algoritma berbasis pohon (*Decision Tree* dan *Random Forest*) sepenuhnya lumpuh dalam mendeteksi kelas minoritas *stroke* (Recall 0% dan 3,6%) [4]. Sebaliknya, model probabilistik *Naive Bayes* terbukti memiliki ketahanan alami (*robustness*) terhadap ketidakseimbangan kelas dengan mempertahankan nilai AUC 0,804 dan Recall 17,3% [7], [11]. Temuan ini memberikan wawasan teoritis baru bagi ilmu data medis bahwa *Naive Bayes* adalah algoritma paling aman sebagai penyaring awal (*screening tool*) pada lingkungan klinis yang belum memiliki sistem penyeimbangan data otomatis.

4. KESIMPULAN

Berdasarkan hasil penelitian yang telah dilakukan, dapat disimpulkan bahwa algoritma *Naive Bayes*, *Random Forest*, dan *Decision Tree* dapat diterapkan untuk memprediksi risiko *stroke* dengan memanfaatkan *Healthcare Stroke Dataset* melalui aplikasi *Orange Data Mining*. Semua tahap penelitian, mulai dari *preprocessing*, pemilihan fitur, pembangunan model, hingga evaluasi dengan metode *10-Fold Cross Validation*, berhasil dilaksanakan dan menghasilkan analisis perbandingan performa ketiga algoritma secara objektif. Evaluasi dilakukan menggunakan berbagai metrik, termasuk *Accuracy*, *Area Under Curve (AUC)*, *Precision*, *Recall*, *F1-Score*, *Matthews Correlation Coefficient (MCC)*, *Confusion Matrix*, dan *Receiver Operating Characteristic (ROC)*. Hasil dari pengujian menunjukkan bahwa algoritma *Decision Tree* mencapai nilai *Accuracy* tertinggi yaitu 0,951, diikuti oleh *Random Forest* yang memiliki nilai 0,948, sementara *Naive Bayes* mendapatkan nilai 0,924. Namun, berdasarkan analisis dari *Confusion Matrix* dan ROC, kedua algoritma tersebut masih mengalami kesulitan dalam mengidentifikasi pasien yang benar-benar terkena *stroke* akibat ketidakseimbangan jumlah data dalam *dataset*. Sebaliknya, algoritma *Naive Bayes* mencatat nilai AUC tertinggi sebesar 0,804, yang menunjukkan kemampuannya yang paling baik dalam membedakan pasien yang mengalami *stroke* dari yang tidak. Secara keseluruhan, *Naive Bayes* menunjukkan performa yang lebih unggul jika dilihat dari berbagai metrik evaluasi, khususnya AUC, MCC, *Confusion Matrix*, dan ROC, sehingga lebih dianggap cocok untuk diaplikasikan pada *dataset* dengan distribusi kelas yang tidak seimbang. Meskipun demikian, setiap algoritma memiliki keunggulan pada metrik tertentu, oleh karena itu pemilihan model sebaiknya tidak hanya didasarkan pada nilai *Accuracy*. Penelitian yang akan datang disarankan untuk menggunakan teknik penyeimbangan data, seperti *SMOTE*, melakukan optimasi parameter model, atau membandingkannya dengan algoritma klasifikasi lainnya agar kemampuan model dalam mendeteksi pasien *stroke* dapat semakin meningkat.

UCAPAN TERIMA KASIH

Penulis mengucapkan rasa syukur yang tulus kepada Tuhan Yang Maha Esa atas semua berkat dan anugerah-Nya, sehingga penelitian dengan judul "Komparasi *Naive Bayes*, *Random Forest*, dan *Decision Tree* pada Prediksi Penyakit *Stroke* Menggunakan *Orange Data Mining*" dapat diselesaikan dengan baik. Selama membuat penelitian ini, penulis mendapat banyak bantuan, petunjuk, dan dukungan dari berbagai orang sehingga semua langkah dalam penelitian bisa berjalan dengan baik. Penulis ingin mengucapkan terima kasih kepada Bapak Jaya Tata Hardinata, M.Kom., yang merupakan dosen pengampu mata kuliah *Machine Learning*. Beliau telah memberikan ilmu, panduan, masukan, serta motivasi selama proses belajar dan dalam penyusunan penelitian ini. Pengetahuan yang diberikan membantu penulis memahami konsep *machine learning*, cara mengolah data, serta bagaimana menerapkan algoritma klasifikasi melalui aplikasi *Orange Data Mining*. Penulis juga ingin menyampaikan rasa terima kasih kepada keluarga, yang telah memberikan dukungan baik dalam doa maupun materi, serta semangat, sehingga penelitian ini dapat terselesaikan dengan lancar. Penulis juga menyampaikan ucapan terima kasih kepada penyedia *dataset Healthcare Stroke Dataset* melalui platform *Kaggle*, karena telah memberikan *dataset* sebagai sumber data utama dalam penelitian ini. Penulis menyadari

bahwa penelitian ini masih memiliki beberapa kekurangan, sehingga sangat diharapkan adanya kritik dan saran yang konstruktif agar penelitian selanjutnya dapat diperbaiki. Semoga penelitian ini bisa bermanfaat dan menjadi acuan untuk pengembangan penelitian di bidang *machine learning*, terutama dalam memprediksi penyakit *stroke*.

REFERENCES

- [1] I. A. I. D. Prayoga and I. G. A. G. A. Kadyanan, "Penggunaan Algoritma C4.5 dan Random Forest guna Meningkatkan Efisiensi Klasifikasi Penyakit Stroke," *Jnatia*, vol. 3, no. 3, pp. 553–564, 2025, [Online]. Available: <https://www.kaggle.com/datasets/fedesoriano/stroke-prediction-dataset>
- [2] A. Ristyawan, A. Nugroho, and T. K. Amarya, "Optimasi Preprocessing Model Random Forest untuk Prediksi Stroke," *JATISI (Jurnal Tek. Inform. dan Sist. Informasi)*, vol. 12, no. 1, 2025, doi: 10.35957/jatisi.v12i1.9587.
- [3] H. Hendriyansyah, A. Irma Purnamasari, and T. Suprpti, "Penerapan Algoritma Decision Tree Dalam Klasifikasi Penyakit Stroke Otak," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 3, pp. 3038–3043, 2024, doi: 10.36040/jati.v8i3.9602.
- [4] N. Biswas, K. M. M. Uddin, S. T. Rikta, and S. K. Dey, "A comparative analysis of machine learning classifiers for stroke prediction: A predictive analytics approach," *Healthc. Anal.*, vol. 2, no. October, p. 100116, 2022, doi: 10.1016/j.health.2022.100116.
- [5] O. Shobayo, O. Zachariah, M. O. Odusami, and B. Ogunleye, "Prediction of Stroke Disease with Demographic and Behavioural Data Using Random Forest Algorithm," *Analytics*, vol. 2, no. 3, pp. 604–617, 2023, doi: 10.3390/analytics2030034.
- [6] S. N. Nasokha and R. A. Firmansyah, "Implementasi Algoritma Naïve Bayes untuk Mendiagnosis Kerentanan Warga terhadap Penyakit Stroke," *Technol. Informatics Insight J.*, vol. 3, no. 2, pp. 114–123, 2024, doi: 10.32639/9hr05d41.
- [7] Wartika and A. Nursikuwagus, "Classification for Diagnosing Stroke Using Orange Data Mining," *J. Math. Sci. Informatics*, vol. 5, no. 1, pp. 83–94, 2025, doi: 10.46754/jmsi.2025.06.007.
- [8] C. Maulana Sidiq, A. Faqih, and G. Dwilestari, "Algoritma Decision Tree C4.5 Digunakan Untuk Mengklasifikasikan Data Stroke," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 2, pp. 1869–1874, 2024, doi: 10.36040/jati.v8i2.8388.
- [9] M. F. Banjar, I. Irawati, F. Umar, and L. N. Hayati, "Analysis of Stroke Classification Using Random Forest Method," *Ilk. J. Ilm.*, vol. 14, no. 3, pp. 186–193, 2022, doi: 10.33096/ilkom.v14i3.1252.186-193.
- [10] Abdul Roni, Maria Fransiska Fitriani, Nazwa Aurellia Ainanur, Sumanto, and Ade Surya Budiman, "Comparison of K-Nearest Neighbor Algorithm Performance and Naïve Bayes in Predicting Stroke Disease," *J. Artif. Intell. Eng. Appl.*, vol. 5, no. 1, pp. 1–6, 2025.
- [11] A. F. Rianny and G. Testiana, "Penerapan Data Mining untuk Klasifikasi Penyakit Jantung Koroner Menggunakan Algoritma Naïve Bayes," *MDP Student Conf.*, vol. 2, no. 1, pp. 297–305, 2023, doi: 10.35957/mdp-sc.v2i1.4388.
- [12] E. Sabna and O. Dewi, "Prediksi Penyakit Stroke menggunakan Algoritma Decision Tree dan Naïve Bayes," *RIGGS J. Artif. Intell. Digit. Bus.*, vol. 4, no. 3, pp. 1294–1299, 2025, doi: 10.31004/riggs.v4i3.2132.
- [13] Q. Lu, S. Li, H. Xu, H. Shen, and Y. Wei, "A hybrid random forest model for stroke risk prediction," *Front. Neurol.*, vol. 17, no. April, 2026, doi: 10.3389/fneur.2026.1822304.
- [14] R. E. Yunus *et al.*, "Stroke Prognostication in Patients Treated with Thrombolysis Using Random Forest," *Open Neuroimag. J.*, vol. 17, no. 1, pp. 1–12, 2024, doi: 10.2174/0118744400298093240520070257.
- [15] Y. Aulia, A. Andriyansyah, S. Suharjito, and S. W. Nensi, "Analisis Prediksi Stroke dengan Membandingkan Tiga Metode Klasifikasi Decision Tree, Naïve Bayes, dan Random Forest," *J. Ilmu Komput. dan Inform.*, vol. 3, no. 2, pp. 89–98, 2024, doi: 10.54082/jiki.90.
- [16] I. Virgiawan, "Analisis Perbandingan Algoritma Naïve Bayes dan Random Forest Dalam Klasifikasi Penyakit Stroke Pada Puskesmas," vol. 6, no. 4, pp. 2807–2814, 2025, doi: 10.47065/bits.v6i4.6771.
- [17] S. Handayani, "Komparasi Metode Klasifikasi Data Mining untuk Prediksi Penyakit Ginjal," *J. Sist. Inf.*, vol. 6, no. 1, pp. 34–40, 2017, doi: 10.51998/jsi.v6i1.118.